

Analyse critique comparative de trois articles sur la positivité de Weil en fenêtre compacte

Évaluation de la fiabilité, examen des erreurs,
et discussion des plateformes de prépublication

18 septembre 2026

Abstract

Ce document propose une analyse critique comparative de trois articles portant sur la positivité de la forme quadratique de Weil en fenêtre compacte : *Weil positivity in compact windows* de Xuefeng Zhu, *The Three Gates* de Marco Desogus, et *Certified Weil Positivity Beyond the Unit Window* de Vincent Liu. L'analyse porte sur la structure des preuves, la transparence méthodologique, la vérifiabilité des certificats, la cohérence interne et la plausibilité mathématique. Une discussion sur la plateforme ALPHAXIV et sur la question du parrainage ARXIV complète l'étude, afin de préciser ce que ces mécanismes garantissent — et ce qu'ils ne garantissent pas — en matière de fiabilité scientifique.

1 Article de Xuefeng Zhu — *Weil positivity in compact windows*

1.1 Nature du résultat

Zhu ne prétend **pas** prouver RH. Il étudie le profil quantitatif $\lambda^*(L) = \inf_f Q(f)/\|f\|_2^2$ sur des fenêtres compactes, et obtient :

- **Bornes inférieures certifiées** (inconditionnelles) : positivité prouvée pour $\text{supp } f \subseteq [-0.8, 0.8]$ avec marge 8.9×10^{-18} , étendant la plage classique $\log 2/2 \approx 0.347$ à 0.8 (facteur 2.3).
- **Bornes supérieures certifiées** (inconditionnelles) : évaluation variationnelle sur le côté géométrique, jusqu'à $L = 2$.
- **Loi empirique de Landau–Widom** : $-\ln \lambda^*(L) \approx 2\pi^2 N(T^*)/\ln N(T^*)$ avec $T^* = 2\pi e^{2L}$.
- **Théorème conditionnel** : sous RH, $\lambda^*(L) \leq \exp(-Le^L)$.
- **Théorème de barrière** (inconditionnel) : tout certificat ponctuel de type « one-stroke » doit atteindre $T_1 = 2\pi e^{A_L}$ avec $A_L \sim 4e^L$.

1.2 Forces

1. **Honnêteté épistémique remarquable.** Zhu distingue systématiquement :

- ce qui est *certifié inconditionnellement* (Théorème 1.4, Théorème 1.6, Tableau 1)
;

- ce qui est *mesuré empiriquement* (loi de Landau–Widom, Conjecture 8.1) ;
- ce qui est *conditionnel* (Théorème 1.7, sous RH) ;
- ce qui est *une barrière* (Théorème 1.8).

Il **rétracte explicitement** une affirmation antérieure (support 2.38) dans la Section 9, avec explication détaillée de l’erreur (confusion entre A_L et A_{eff} , mauvais sens de l’inégalité). C’est un signe de sérieux.

2. **Mécanisme de preuve clair et vérifiable.** La « one-stroke reduction » (Théorème 3.1) est une réduction en dimension finie via une enveloppe ponctuelle sur le symbole de Weil. Les ingrédients sont élémentaires (Binet, Bernstein, Cholesky vérifié), et le budget d’erreur se referme avec 24 ordres de grandeur de marge.
3. **Le lemme d’optimalité de la constante du peigne** (Lemme 3.2) est **correct et important** : $\sup_t P_L(t) = A_L$ exactement, par équidistribution de Weyl. Cela justifie rigoureusement que la barrière est intrinsèque à la classe des certificats ponctuels.
4. **Distinction soigneuse** entre le côté géométrique (calcul inconditionnel) et le côté zéros (moteur de proposition et validation croisée). La Remarque 2.2 (positivity vs cancellation) est un point méthodologique crucial que beaucoup d’auteurs négligent.
5. **Parité** : le Théorème 8.1 (ground state simple et pair) est un résultat certifié, avec séparation bilatérale, qui répond à une hypothèse spectrale explicitement demandée par Connes–Consani–Moscovici et Connes–van Suijlekom.

1.3 Faiblesses et points de vigilance

1. **Le résultat principal (Théorème 1.4) est modeste en portée.** Support 1.6 vs $\log 2 \approx 0.693$ classique : un facteur 2.3, pas un saut qualitatif. La borne 8.9×10^{-18} est minuscule. Zhu le reconnaît, mais cela signifie que la méthode ne « passe pas l’échelle ».
2. **La loi de Landau–Widom est ajustée, non dérivée.** Le fait que $2\pi^2$ soit « familier » ne constitue pas une preuve. Zhu le dit explicitement (Remarque 8.2 : « la constante est ajustée »), mais un lecteur pressé pourrait confondre le fit avec un théorème. La discrimination de modèles par test hors échantillon est astucieuse mais repose sur seulement 4 points.
3. **Le Théorème 10.1 (conditionnel)** donne $\exp(-Le^L)$, alors que la loi empirique suggère $\exp(-2\pi^2 N(T^*)/\ln N(T^*))$ avec $T^* = 2\pi e^{2L}$, soit une décroissance beaucoup plus rapide. L’écart entre e^L et e^{2L} dans l’exposant est reconnu (Remarque 10.2 : obstruction at full saturation), mais cela montre que le théorème conditionnel ne capture qu’une version très affaiblie de la loi.
4. **Dépendance à des données de zéros de haute précision** pour le pipeline supérieur. Zhu le note (Remarque 7.3) : les zéros sont des « calculs de haute confiance », pas des enclosures formellement vérifiées. Pour les bornes **supérieures** (variationnelles), cela n’affecte pas la validité (n’importe quel vecteur d’essai donne une borne valide). Mais cela signifie que le pipeline du côté des zéros n’est pas un certificat au sens strict.
5. **Le Théorème de barrière (1.8) est inconditionnel mais limité** : il ne s’applique qu’aux certificats à enveloppe ponctuelle. Zhu le dit (Remarque 4.2) : « converting it into a proved lower bound on the cost of an arbitrary positivity certificate is open ».

1.4 Évaluation globale de Zhu

Fiabilité : très élevée. Le papier est structuré, honnête, vérifiable, avec une distinction claire entre prouvé / mesuré / conjecturé / conditionnel. La rétractation explicite est un gage de sérieux. Les résultats sont modestes mais solides. La valeur principale est **méthodologique** : quantifier la difficulté de RH via $\lambda^*(L)$, et cartographier la barrière.

Verdict : article fiable, prudent, sans surendevancement.

2 Article de Marco Desogus — *The Three Gates*

2.1 Nature du résultat

Desogus **prétend prouver RH** (Théorème 1.1 : $A_{a,-} \succeq 0$ pour tout $a > 0$, donc RH). L'argument est structuré en trois « portes » :

- **Porte I** : théorème de covariance cellulaire (convexité stricte de $\mu''_{Y,d}$).
- **Porte II** : routage arithmétique exact, court-circuit de Feshbach, transport Cauchy–Carleman, graphe harmonique, « MASTER » scalaire.
- **Porte III** : fermeture et critère de Weil impair restreint.

Il utilise un cadre très élaboré : opérateurs enracinés, Schur complements, court-circuits variationnels, congruences, certificats par intervalles, bases de Legendre, etc.

2.2 Forces apparentes

1. **Architecture ambitieuse.** L'auteur construit une machinerie complexe avec de nombreux lemmes techniques, une chaîne de dépendances explicite, et une tentative de traiter le canal impair complet.
2. **Quelques résultats intermédiaires plausibles.** Le Théorème 2.2 (critère de Weil impair restreint) est un énoncé classique correct : RH équivaut à la positivité sur le cœur test impair réel. La preuve esquissée (construction d'un test impair isolant une orbite de zéro hors axe) est dans l'esprit des arguments standards.
3. **Tentative de certification** : l'auteur parle de certificats par intervalles, de Cholesky vérifié, de constantes explicites. C'est louable en intention.

2.3 Faiblesses graves et signaux d'alerte

C'est ici que l'analyse devient critique. Il y a **plusieurs problèmes structurels majeurs** :

2.3.1 (a) Le résultat est extraordinairement fort, mais la preuve ne l'est pas

Prouver RH est l'un des problèmes les plus difficiles des mathématiques. Un article qui prétend le faire doit avoir une preuve **entièrement rigoureuse, vérifiable, sans trou**. Or :

- La preuve repose sur une **chaîne de dizaines de lemmes techniques** (Portes I, II, III, Appendices A–F), chacun avec des constantes explicites, des congruences, des certificats numériques.

- **Aucun de ces lemmes n'est trivial**, et leur combinaison devrait être vérifiée par une communauté d'experts sur plusieurs mois, voire années.
- L'article est **auto-publié** (pas de revue par les pairs, pas de soumission à une revue).

2.3.2 (b) Confusion entre modèle scalaire et opérateur

Le Point d'audit 3.1 (Porte I) est révélateur : l'auteur **admet lui-même** que le Corollaire 3.3 (« $\mathcal{P}_{Y,d}(s) > 0$ ») est un théorème sur un **modèle scalaire hybride/coquille**, et qu'il **n'identifie pas** $\mathcal{P}_{Y,d}$ à un pivot de Schur de l'opérateur de Weil localisé. Il dit que « le problème de transfert à valeurs opératorielles est gardé séparé dans la Porte II ».

C'est un **aveu de gap** : la Porte I prouve quelque chose sur un modèle, pas sur l'opérateur. Il faut ensuite un « lift » opératoriel (que l'auteur appelle AWGC-OP, MASTER-OP, etc.). Ce lift est précisément le point le plus fragile.

2.3.3 (c) Prolifération de constantes et de lemmes « ajustés »

L'article pullule de constantes numériques : $439/250$, $189/250$, $9617/10400$, $63183/10400$, $q_F = 0.9922\dots$, $\gamma_7 = 0.4533\dots$, $d_* = 0.4934\dots$, $\delta_* = 0.4515\dots$, $\beta = 0.0484\dots$, etc. Chacune est justifiée par un lemme, mais :

- Ces constantes sont souvent **choisies pour que les inégalités passent** (par exemple $439/250 = 1 + 189/250$, où $189/250$ est « le coefficient mixte établi ci-dessous »).
- Il y a un **risque élevé d'erreur cumulative** dans une chaîne aussi longue.
- Les certificats numériques (Cholesky, intervalles) ne sont pas fournis dans le texte, et l'auteur renvoie à un « paquet supplémentaire » non public.

2.3.4 (d) Le « MASTER » et les lemmes P2/P3 sont opaques

Les Lemmes 13.x (MASTER-P2, MASTER-P3, FOLD-FORCING, etc.) sont présentés comme des identités exactes, mais :

- Leur preuve repose sur des hypothèses de compatibilité (par exemple, « after the full-form transport, fixed-target gauge and endpoint Möbius fold used in the common-cut argument », Corollaire 13.6).
- Ces hypothèses ne sont pas vérifiées indépendamment ; elles sont **postulées** comme découlant de constructions antérieures.
- L'audit 11.8 (« Compatibility limit of the root-adapted normalization ») **admet** que la normalisation adaptée à la racine ne commute pas avec la transformée de Fourier ponctuelle, et qu'« une estimation opératorielle supplémentaire serait nécessaire ». L'auteur dit que cette compatibilité est « résolue » par la restriction au graphe A_k -harmonique, mais cela reste un point de fragilité.

2.3.5 (e) Le « certificat de base $Y = 7$ » est invoqué, pas reproduit

Le Certificat 13.x (« Arithmetic base at $Y = 7$ ») dit : « At $a_7 = \frac{1}{2} \log 7$, the rigorous Arb/Rump fifth-cell certificate for the same odd localized core proves strict positivity even in the stronger normalization $C_{a_7,-} - 4|s^{\text{ph}}\rangle\langle s^{\text{ph}}| \succ 0$. Its smallest certified eigenvalue lower endpoint is $1.0705 \times 10^{-24} > 0$. »

Mais :

- Ce certificat n'est **pas détaillé** dans l'article (contrairement à l'Appendice F qui prétend le faire, mais la page 12 du PDF est corrompue et illisible).
- L'auteur renvoie à un « paquet supplémentaire » non fourni.
- La valeur 1.07×10^{-24} est **extraordinairement petite** : la matrice est à peine définie positive. Une erreur numérique ou une constante mal placée pourrait facilement inverser le signe.

2.3.6 (f) Absence de vérification par les pairs et de reproduction indépendante

L'article est **auto-publié**, sans revue par les pairs, sans reproduction indépendante documentée. L'auteur mentionne un « programme de reproduction » mais ne fournit pas de lien public ni de hash vérifiable. C'est un **problème majeur** pour un résultat de cette ampleur.

2.3.7 (g) Le style « fantasy interludes » est un signal d'alarme

Les interludes narratifs (Herr G.F.B.R., Marc, Johnny Nash, Dama Noether, etc.) sont présentés comme « purely expository » et « affectionate fictional homages ». Mais dans un article prétendant prouver RH :

- Cela **dilue le sérieux** et rend la lecture plus difficile.
- Cela **masque** la structure logique derrière des métaphores (« Doppelganger », « Beholder », « Wizard »).
- Cela pourrait être un **indice de génération assistée par IA** (l'auteur admet l'usage de GPT-5.6 Sol dans les remerciements).

2.3.8 (h) Le théorème principal est « trop beau pour être vrai »

Prouver RH nécessiterait de résoudre :

- La conjecture de Hilbert–Pólya.
- Le problème de l'universalité de la positivité de Weil.
- Le contrôle de l'opérateur non borné sur tout support.
- Le comportement asymptotique des zéros.

Aucune de ces difficultés n'est résolue par une « réduction en trois portes » avec des certificats numériques à 10^{-24} près. Les tentatives sérieuses de prouver RH (Connes, Bombieri, etc.) reconnaissent explicitement la difficulté et ne prétendent pas avoir une preuve complète.

2.4 Évaluation globale de Desogus

Fiabilité : très faible. L'article prétend prouver RH, mais :

- La preuve repose sur une chaîne de lemmes non vérifiés indépendamment.
- Les certificats numériques ne sont pas fournis.
- Les points de fragilité sont admis par l'auteur lui-même (audits).

- L’usage d’IA générative est reconnu, mais l’auteur ne distingue pas clairement ce qui est de lui et ce qui est généré.
- Le style narratif masque les gaps.
- Aucune revue par les pairs.

Verdict : article non fiable. À traiter avec une extrême prudence. Probablement contient des erreurs, peut-être fondamentales.

3 Article de Vincent Liu — *Certified Weil Positivity Beyond the Unit Window*

3.1 Nature du résultat

Liu ne prétend **pas** prouver RH (il le dit explicitement : « Fixed-window positivity does not prove the Riemann Hypothesis »). Il prouve :

- **Théorème A** : positivité sur la fenêtre unité $[-1, 1]$ avec marge 2^{-151} .
- **Théorème B** : positivité sur la fenêtre $[-17/16, 17/16]$ avec marge 2^{-49162} , incluant la puissance de nombre premier 8.
- **Théorème d’Obstruction** : une localisation spécifique (tail-dropping) a une limite haute fréquence négative si $L > \log 8/2$, et aucune correction compacte fixée ne la répare.
- **Théorème de Certification** : transfert d’une inégalité finie à erreur payée vers l’espace de Hilbert complet, avec fermeture exacte de Gram.

3.2 Forces

1. Transparence et rigueur méthodologiques exceptionnelles.

- Liu distingue clairement les **entrées analytiques** (Appendices A–E) des **calculs finis** (certificats entiers).
- Il fournit un **protocole de reproduction** détaillé (Appendice F) avec hashes SHA-256, commandes exactes, temps d’exécution, et vérification par deux compilateurs indépendants.
- Il **admet les limites** : le Théorème d’Obstruction ne donne pas un test de Weil négatif, et ne s’applique qu’à une stratégie spécifique.

2. Séparation claire des ordres d’erreur. Le Théorème de Certification (Section 6) sépare l’erreur croisée de projection $O(r)$ de l’erreur de complément $O(r^2)$. C’est une **amélioration technique réelle** par rapport aux approches qui traitent toutes les erreurs comme une seule perturbation.

3. Certificats entiers exacts. Le point clé : la positivité est réduite à des **inégalités entières exactes** (équation (12) et appendice E), vérifiables par arithmétique exacte. Pas de pivot flottant comme certificat de signe. C’est un standard de rigueur élevé.

4. **Obstruction analytique indépendante.** Le Théorème d'Obstruction (Section 5) est un résultat **négatif** important : il montre qu'une stratégie naturelle (tail-dropping) échoue au-delà de $L = \log 8/2$. C'est un résultat falsifiable et utile pour la communauté.
5. **Distinction entre analyse et computation.** Liu insiste : les arguments analytiques (Appendices A–E) et les vérifications arithmétiques (Appendice F) sont **tous deux nécessaires**. Il ne substitue pas l'un à l'autre.
6. **Déclaration IA détaillée et honnête.** Liu admet l'usage de ChatGPT et Codex, mais précise que les suggestions mathématiques générées n'ont **pas** été traitées comme autorité, et que chaque affirmation a été vérifiée par argument explicite ou certificat. C'est un standard de transparence élevé.

3.3 Faiblesses et points de vigilance

1. **Portée limitée.** Liu prouve la positivité sur $[-17/16, 17/16]$, soit $L \approx 1.0625$. C'est **plus grand** que Zhu ($L = 0.8$ pour la borne inférieure certifiée). Mais sa marge est **beaucoup plus petite** ($2^{-49162} \approx 10^{-14798}$ vs 8.9×10^{-18}). Ce sont des résultats **complémentaires**, pas directement comparables.
2. **Le Théorème d'Obstruction est spécifique.** Il ne s'applique qu'à une stratégie de « tail-dropping » particulière. D'autres stratégies (non compactes, ou avec des queues différentes) ne sont pas exclues. Liu le dit explicitement.
3. **Complexité technique élevée.** Les Appendices B–E sont denses : estimations de Bessel, quadrature de Gauss, congruences entières, constantes explicites. La vérification indépendante demanderait un effort considérable.
4. **Dépendance à des faits classiques externes.** Liu utilise des résultats de NIST DLMF, Trudgian, Schoenfeld, Dusart, etc. C'est normal, mais cela signifie que la preuve n'est pas totalement auto-contenue.
5. **Pas de formalisation par assistant de preuve.** Liu le note (Section 7) : « A formal proof-assistant treatment has not been carried out ». C'est une limite, mais pas une faiblesse en soi (la plupart des articles de mathématiques ne sont pas formalisés).

3.4 Évaluation globale de Liu

Fiabilité : très élevée. L'article est :

- **Précis** : constantes explicites, bornes d'erreur séparées, certificats entiers.
- **Honnête** : limites admises, portée claire, pas de surrevendication.
- **Reproductible** : protocole détaillé, hashes, deux compilateurs indépendants.
- **Techniquement solide** : séparation des ordres d'erreur, obstruction analytique, fermeture exacte de Gram.

Verdict : article fiable, rigoureux, avec une portée modeste mais des garanties de certification élevées.

Critère	Zhu	Desogus	Liu
Prétend prouver RH ?	Non	Oui	Non
Portée	$L = 0.8$ (inf.), $L = 2$ (sup.)	Toute fenêtre	$L = 1.0625$
Marge certifiée	8.9×10^{-18}	Non vérifiable	2^{-49162}
Inconditionnel ?	Oui (inf. et sup.)	Prétendu	Oui
Certificats fournis ?	Partiellement	Non	Oui (protocole détaillé)
Revue par les pairs ?	Non	Non	Non
Usage d'IA déclaré ?	Non mentionné	Oui (GPT-5.6 Sol)	Oui (ChatGPT, Codex)
Style	Académique sobre	Narratif (fantasy)	Académique détaillé
Distinction prouvé/mesuré/conjecturé	Excellente	Faible	Excellente
Reproductibilité	Partielle	Faible	Élevée
Falsifiabilité	Élevée	Faible	Élevée
Cohérence interne	Bonne	Problématique	Bonne
Plausibilité	Élevée	Faible	Élevée

Table 1: Synthèse comparative des trois articles.

4 Synthèse comparative

5 Classement de fiabilité

5.1 1^{er} : Vincent Liu — le plus fiable

- Résultats modestes mais **rigoureusement certifiés**.
- Protocole de reproduction **détaillé et vérifiable**.
- Distinction claire entre analyse et computation.
- Obstruction analytique **indépendante**.
- Transparence sur l'IA.
- **Recommandation : article de référence pour la certification en fenêtre fixée.**

5.2 2^e : Xuefeng Zhu — fiable, mais portée plus large et moins certifiée

- Résultats **plus ambitieux** (loi asymptotique, barrière, parité).
- Distinction claire entre prouvé / mesuré / conjecturé / conditionnel.
- **Rétractation explicite** d'une erreur antérieure : gage de sérieux.
- Mais certificats numériques **moins détaillés** que Liu.
- **Recommandation : article de référence pour la compréhension quantitative de la difficulté de RH.**

5.3 3^e : Marco Desogus — non fiable

- Prétend prouver RH, mais :
 - Chaîne de lemmes non vérifiée indépendamment.
 - Certificats numériques non fournis.
 - Points de fragilité admis par l’auteur (audits).
 - Style narratif masquant les gaps.
 - Aucune revue par les pairs.
 - Usage d’IA générative reconnu mais mal délimité.
- **Recommandation : à ne pas considérer comme une preuve de RH. À lire avec une extrême prudence, voire à ignorer.**

6 Conclusion

Les trois articles traitent du même sujet (positivité de Weil en fenêtre compacte), mais avec des **postures épistémiques radicalement différentes** :

- **Liu** : « Voici ce que je certifie, voici comment le vérifier, voici les limites. » → **Science normale.**
- **Zhu** : « Voici ce que je mesure, voici ce que je conjecture, voici pourquoi la route est bloquée. » → **Science normale avec ambition.**
- **Desogus** : « Voici la preuve de RH, en trois portes. » → **Surrevendication, probablement erronée.**

Si vous devez n’en retenir qu’un, je recommande **Liu** pour la fiabilité des certificats, et **Zhu** pour la compréhension quantitative du problème. **Desogus** ne devrait pas être utilisé comme référence sans une vérification indépendante exhaustive, qui n’a manifestement pas été effectuée.

Un point important : **aucun des trois articles ne prouve RH**. Liu et Zhu le disent explicitement ; Desogus prétend le contraire, mais sa prétention n’est pas soutenue par une preuve vérifiable.

7 Discussion : alphaXiv, arXiv et le parrainage

7.1 Qu’est-ce qu’alphaXiv ?

ALPHAXIV a été fondé vers 2023 par deux étudiants de Stanford. C’est essentiellement une **plateforme de discussion ouverte basée sur arXiv**. Son mécanisme central est le suivant :

- il suffit de remplacer `arxiv.org` par `alphaxiv.org` dans l’URL d’un article arXiv pour ouvrir sa page de discussion ;
- à gauche s’affiche le texte original, à droite les commentaires ligne par ligne et les réponses, auxquelles les auteurs peuvent participer ;

- la plateforme propose des résumés assistés par IA, une recherche plein texte, un mode « Deep Research » pour la revue de littérature ;
- une extension Chrome signale sur les pages arXiv la présence de discussions associées.

Point clé : **la source des articles d’alphaXiv, c’est arXiv lui-même**. Elle ne reçoit pas de soumissions indépendantes et ne pratique aucune évaluation scientifique autonome.

7.2 Le problème de la « fiabilité » d’alphaXiv

7.2.1 Elle n’apporte aucune garantie supplémentaire

Une équipe de bibliothécaires a explicitement évalué alphaXiv : sa base de données repose sur des **préprints**, c’est-à-dire des manuscrits **non encore soumis à l’évaluation par les pairs**. arXiv lui-même procède à un **filtrage préliminaire** assuré par des experts bénévoles, mais cela **n’équivaut pas à une évaluation par les pairs**.

alphaXiv ajoute à cette base une **fonction de discussion communautaire**, mais les commentaires ne constituent **pas une évaluation formelle** et ne garantissent pas la compétence des commentateurs.

7.2.2 Elle résout la « découvrabilité », pas la « crédibilité »

La valeur d’alphaXiv réside dans :

- rendre les articles arXiv **plus faciles à trouver et à discuter** ;
- fournir des outils de **lecture et de compréhension** assistés par IA ;
- permettre aux auteurs et lecteurs d’**interagir directement**.

Ce qu’elle **ne résout pas** :

- la correction mathématique d’un article ;
- l’originalité et la priorité d’un article ;
- la présence ou non de « contenu poubelle » généré par IA.

7.3 Le parrainage arXiv : ce qu’il garantit et ce qu’il ne garantit pas

La politique officielle d’arXiv précise que l’endosseur doit **connaître personnellement le candidat ou avoir vu l’article**, et qu’il ne doit endosser que des personnes dont il est **confiant qu’elles soumettront des articles de haute qualité**. Depuis janvier 2026, cette exigence s’applique à **tous les domaines d’arXiv**. Le taux de rejet d’arXiv est passé d’environ 4 % à **10–12 %** avec l’arrivée des soumissions générées par IA.

Donc oui : **le parrainage arXiv constitue un filtre humain minimal**. Mais arXiv le dit elle-même, très explicitement : « **Material is not peer-reviewed by arXiv** — the contents of arXiv submissions are wholly the responsibility of the submitter and are presented “as is” without any warranty or guarantee. » Et la page officielle sur le parrainage précise : « **The endorsement process is not peer review**. We do not expect you to read the paper in detail, or verify that the work is correct. »

Autrement dit, l’endosseur **ne certifie pas la validité mathématique** de l’article. Il certifie seulement que l’auteur est probablement un chercheur légitime.

7.4 Ce que alphaXiv garantit : rien de plus

ALPHAXIV n'a **aucun filtre** à l'upload : n'importe qui peut déposer n'importe quel PDF. La plateforme a été conçue comme une **couche de discussion et d'annotation** sur arXiv. Elle ne prétend **pas** offrir de validation scientifique. Son propre modèle repose sur les **signaux sociaux** (vues, commentaires, partages) pour trier le contenu, pas sur un filtre éditorial.

7.5 Le ressenti initial est-il justifié ?

Partiellement oui, mais avec une nuance importante.

Sur le plan du **filtrage humain minimal** : un article qui a passé le parrainage arXiv a bénéficié d'un regard humain qui a dit « cette personne est probablement légitime ». Un PDF uploadé sur alphaXiv n'a bénéficié d'**aucun** tel regard. En ce sens très précis, oui, **le contenu arXiv est « un tout petit peu moins en balance »** qu'un contenu alphaXiv, parce qu'il y a eu un filtre — même minimal — à l'entrée.

Mais sur le plan de la **validité scientifique** : ni l'un ni l'autre ne garantit quoi que ce soit. Le parrainage arXiv ne vérifie **pas** la correction des mathématiques. Un article arXiv peut être **totalelement faux** malgré le parrainage. Un article alphaXiv peut être **totalelement correct** malgré l'absence de filtre.

7.6 Conséquence pour l'évaluation des trois articles

Cela ne change pas l'analyse. Liu (sur alphaXiv seulement) produit des **certificats transparents et vérifiables** — c'est le contenu qui fait sa fiabilité, pas la plateforme. Desogus (sur arXiv, donc parrainé) produit une **preuve de RH non vérifiable et auto-contradictoire** — le parrainage ne le sauve pas. Le **contenu**, encore une fois, est le seul critère qui compte vraiment.

7.7 Note personnelle sur la situation des chercheurs indépendants

La politique de parrainage d'arXiv, bien que destinée à filtrer le « AI slop », a un effet secondaire : elle **exclut de facto les chercheurs indépendants** qui n'ont ni adresse institutionnelle, ni collègue arXiv disposé à les endosser. C'est une tension réelle entre la lutte contre le contenu de mauvaise qualité et l'ouverture scientifique. Les plateformes comme ALPHAXIV, ZENODO, HAL ou OSF offrent des alternatives, mais aucune n'a la visibilité d'arXiv. Pour un chercheur indépendant, la stratégie la plus robuste reste de **produire un contenu auto-certifiant** — preuves explicites, certificats reproductibles, protocoles détaillés — car c'est le contenu, et non la plateforme, qui finit par établir la crédibilité.